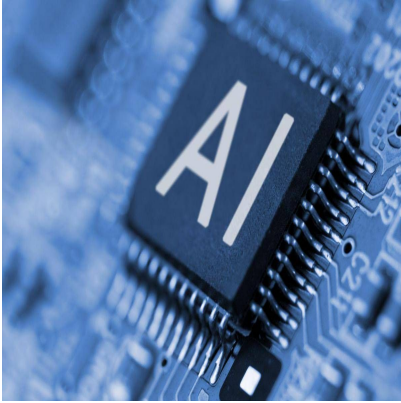




Intelligenza Artificiale pu² ucciderci tutti in 10 anni[•], l'allarme del ricercatore

Descrizione

(Adnkronos) [•]

Intelligenza Artificiale potrebbe uccidere tutti gli esseri umani[•] nel prossimo decennio. E[•] l'allarme che Evan Hubinger, un ricercatore di punta di Anthropic, uno dei colossi dell'AI, fa scattare con un post su X. [•] Crediamo davvero che l'AI possa uccidere tutti gli esseri umani. Personalmente ritengo che la probabilit¹ sia superiore al 10% nel prossimo decennio[•], scrive in un post che viene rilanciato dall'emittente Cbs. Hubinger si presenta come ricercatore dell'azienda con sede a San Francisco. [•] Credo che Anthropic stia facendo del suo meglio, ma non abbiamo ancora un piano per risolvere il problema dell'allineamento per la superintelligenza e non siamo chiaramente sulla strada giusta per riuscirci. Penso che il rischio dai modelli attuali sia basso. Quello che mi preoccupa [•] la superintelligenza derivante dal miglioramento autonomo: come abbiamo detto, sta accadendo pi¹ velocemente di quanto pensassimo[•], aggiunge commentando il post di un collega che ha appena presentato le proprie dimissioni.

Oggi mi sono dimesso da Anthropic. Ho trascorso gli ultimi tre anni a occuparmi di ricerca sul pre-addestramento sia presso OpenAI che presso Anthropic. Nessuna delle due aziende sta agendo in modo responsabile[•], le parole di Jacob Coxon in un post su X. [•] Stanno correndo a tutta velocit¹ verso una superintelligenza capace di auto-migliorarsi, mettendo a repentaglio le nostre vite. In OpenAI, molti non hanno interiorizzato totalmente la portata della posta in gioco per la civilt¹. In Anthropic, la posta in gioco [•] ben compresa, ma l'azienda [•] incastrata in una corsa per arrivare prima degli altri: ritengono che nessun altro agir¹ responsabilmente e che quindi[•] non [•] debbano farlo nemmeno loro, nonostante i rischi[•].

La Cbs ha chiamato in causa l'AI Security Institute, l'organizzazione internazionale creata dai governi e incaricata di analizzare, testare e ridurre i rischi legati ai sistemi di intelligenza artificiale pi¹ avanzati. [•] L'AI Security Institute continua a collaborare strettamente con i partner del settore, tra cui Anthropic, per rendere i modelli pi¹ sicuri[•], le parole di un portavoce del Cabinet Office britannico alla Cbs, sottolineando che [•] proprio la scorsa settimana[•] [•] stato testato [•] GPT-6 Astra[•], il

modello â??piÃ¹ potenteâ?• di OpenAI, prima del suo rilascio al pubblico. â??I rischi non si fermano ai confini nazionali e nessun Paese puÃ² affrontarli da solo. Il Regno Unito continuerÃ a testare i modelli piÃ¹ avanzati, a sviluppare una rigorosa comprensione scientifica delle loro capacitÃ e dei relativi rischi, e a garantire che le decisioni politiche si basino su evidenze concreteâ?•, ha affermato il portavoce.

â??

internazionale/esteri

webinfo@adnkronos.com (Web Info)

Categoria

1. Comunicati

Tag

1. Ultimora

Data di creazione

Settembre 9, 2026

Autore

redazione

default watermark